

Section 7 — Case Study E (Empirical)

Identity Compromise of the AI Data Estate: The Snowflake Customer Credential Campaign (2024)

Drop-in section for the AISIF Practitioner White Paper · Formatted to the Section 7 case-study convention · Prepared July 2026

Case Study E extends the empirical series begun with Case Study D. Its inclusion in an AI security framework requires explicit justification, which is itself the analytical point: the campaign was not an attack on an AI system, yet it compromised the class of platform — the enterprise cloud data estate — from which AI training sets, feature stores and retrieval corpora are drawn. Where Case Study D demonstrated that AI systems fail through conventional infrastructure, Case Study E demonstrates that AI Programs inherit, at full severity, the identity risk of their data supply chain. AISIF treats this estate as in-scope for the Data layer; a framework that excludes it would certify AI pipelines as secure while their source data is exfiltrated upstream.

Case Study Card E — presented in the standard Section 7 card format.

Case Study E — Identity Compromise of the AI Data Estate: The Snowflake Customer Campaign Cloud Data Platform Customers (empirical, disclosed May-June 2024)	
Compromise vector	Financially motivated threat actor UNC5537 systematically accessed approximately 165 Snowflake customer tenants using valid customer credentials harvested by commodity infostealer malware (Vidar, LummaC2, RISEPRO, Redline and others), with the earliest harvested credential dating to November 2020. Access succeeded because the affected accounts lacked multi-factor authentication, credentials had never been rotated after the original infections, and tenants had no network allow-lists restricting access to trusted locations. Many infostealer infections resided on third-party contractor endpoints also used for personal activity, frequently holding elevated privileges across multiple client environments. Data was exfiltrated at scale and victims extorted; confirmed victims included Ticketmaster/Live Nation, Santander, AT&T, Advance Auto Parts and Neiman Marcus. Snowflake's own platform was not breached.
Risk-type refinement	Commonly labelled 'supply chain compromise and token reuse'. The public record supports a more precise characterisation: identity-centric SaaS compromise under a shared-responsibility failure. The supply-chain element is real but specific — infostealer-infected contractor endpoints and SaaS concentration risk — while the mechanism was stale, single-factor passwords rather than session-token replay. Precision matters for control selection: token-lifecycle controls alone would not have prevented this campaign; identity-trust enforcement would have.
AISIF layers implicated	Data Layer (primary asset at risk) — the compromised tenants constitute the analytical data estate that anchors enterprise AI: the corpora from which training sets, feature stores and retrieval indexes are drawn. Exfiltration of this estate is exfiltration of AI-relevant data at source. Operations Layer (primary failure locus, customer side) — single-factor authentication on internet-reachable tenants; no credential rotation over

	multi-year horizons; no network policy; no anomaly detection on authentication or bulk-egress behaviour. Governance Layer (root cause) — shared-responsibility ambiguity: platform-available controls (MFA, network policy) were not mandated by customer governance; third-party contractor access was not held to the customer's own endpoint and identity standards. Application Layer (secondary) — service accounts and pipeline credentials connecting analytics and AI workloads to the data estate share the same identity weaknesses; a harvested pipeline credential is indistinguishable from a harvested human credential. Model Layer — not implicated directly; latent downstream risk only (exfiltrated or manipulated source data propagating into training and retrieval corpora).
Failure mode	Shared-responsibility governance failure manifesting as identity-trust absence at the Operations layer, with third-party endpoint compromise as the initial access vector. No vulnerability in the platform and no novel technique was required: the campaign's scale was a consequence of the infostealer credential marketplace meeting unmandated identity controls across a concentrated SaaS estate.
What single-framework analysis missed	AI-specific frameworks are silent here: no OWASP LLM category, no ATLAS technique, no AI RMF action addresses single-factor authentication on a data warehouse. Conversely, traditional identity guidance existed but was not binding: the platform offered MFA and network policies without enforcing them, and customer standards did not close the gap. The failure sits in the seam between provider capability and customer obligation — a seam that only a shared-responsibility control mapping, of the kind the CSA AICM formalises for AI roles, makes visible and assignable.
AI-SIF mitigation	Governance: bind platform-available security controls to mandatory customer standards; extend identity and endpoint requirements to third-party contractors as a condition of access; inventory the data estate feeding AI workloads. Data: classify and monitor the analytical estate as AI-relevant crown-jewel data; alert on anomalous bulk egress. Application: dedicated least-privilege service identities for AI/analytics pipelines; short-lived, vaulted credentials; no shared human-service credentials. Operations: phishing-resistant MFA on all tenants; credential rotation with defined maximum lifetimes and revocation on infostealer intelligence; network allow-lists; authentication anomaly detection integrated into the recursive feedback loop.
Standards activated	ISO/IEC 27001:2022 Annex A (identity and access management, supplier relationships, monitoring) · NIST SP 800-207 Zero Trust Architecture · CSA AI Controls Matrix shared-responsibility model (AI Customer role) · NIST AI RMF Govern (third-party risk) · MITRE ATT&CK Valid Accounts (T1078) — noting ATT&CK, not ATLAS, carries the relevant technique.

7.5.1 Incident Overview

Between April and June 2024, a financially motivated threat actor tracked by Mandiant as UNC5537 conducted one of the largest credential-based cloud intrusion campaigns on record, accessing approximately 165 Snowflake customer tenants and exfiltrating data for sale and extortion. The campaign became public in late May 2024 when 560 million Ticketmaster customer records were advertised on a criminal forum; Snowflake, Mandiant and CrowdStrike issued coordinated disclosures shortly afterwards, and Snowflake published hardening guidance on 30 May 2024. Confirmed or publicly named victims included Ticketmaster/Live Nation, Santander, AT&T, Advance Auto Parts and Neiman Marcus. The principal alleged actor was arrested in Canada in October 2024.

The joint investigation established three findings that define the case. First, Snowflake's own enterprise environment was not breached; every intrusion traced to valid customer credentials. Second, those credentials were harvested by commodity infostealer malware — families including Vidar, LummaC2, RISEPRO and Redline — from infected endpoints, with the earliest observed infection dating to November 2020; the credentials had circulated on criminal marketplaces for years and had never been rotated. Third, the compromised tenants shared a common configuration profile: no multi-factor authentication, no network allow-lists, and passwords unchanged since the original infections. Mandiant characterized the campaign as requiring no novel or sophisticated technique; its scale was the product of the infostealer marketplace meeting unmandated identity controls. A significant fraction of the infected endpoints belonged to third-party contractors using personal or unmonitored machines, often carrying elevated privileges across multiple client environments.

For AISIF, the incident is the canonical shared-responsibility case. Every preventing control existed and was available on the platform; none was mandatory, and customer governance did not make it so. The failure sits in the seam between provider capability and customer obligation — precisely the seam that role-based shared-responsibility mappings exist to close.

7.5.2 Layer-by-Layer Analysis

Table E-1. AISIF layer implication analysis.

AISIF Layer	Implication	Analysis
Governance	Root cause	Two governance failures compound. Shared-responsibility ambiguity: the platform offered MFA and network policies but did not enforce them, and customer standards did not mandate their use — leaving the estate's identity posture to default configuration. Third-party governance: contractor access was not bound to the customer's own endpoint and identity requirements, so personal, infostealer-infected machines held privileged credentials into the data estate. Neither failure is visible to an assessment scoped to the AI application alone.
Data	Primary asset at risk	The compromised tenants are the analytical data estate — the upstream source of training corpora, feature stores and retrieval indexes for enterprise AI. The campaign demonstrates that this estate carries crown-jewel sensitivity independent of any AI application built upon it, and that its compromise propagates silently into every downstream AI workload that consumes it. Classification, egress monitoring and inventory of AI-feeding

		datasets are Data-layer obligations that would have both constrained and detected bulk exfiltration.
Model	Not attacked; latent	No model was targeted. The latent path is nonetheless material for the framework: an adversary with write access to source data holds a poisoning capability against every model later trained or grounded on that estate (the Case Study B mechanism, acquired through the Case Study E vector), and exfiltrated corpora may expose data intended for proprietary model training.
Application	Secondary	AI and analytics pipelines authenticate to the data estate with service accounts and pipeline credentials subject to the same identity weaknesses exploited here. A harvested service credential is operationally indistinguishable from a harvested human credential — and frequently longer-lived, broader in privilege and less monitored. The campaign's mechanics apply to non-human identities without modification.
Operations	Primary failure locus	All proximate causes are identity-trust absences on the customer side: single-factor authentication on internet-reachable tenants; credentials unrotated across multi-year horizons despite circulating in criminal marketplaces; no trusted-location network policy; and no behavioural detection on authentication or bulk egress — intrusions were identified through external threat intelligence and victim extortion, not internal telemetry.

7.5.3 Control Domain Mapping

Table E-2 maps observed failures to AISIF control domains and operative framework references, following the Table 3 crosswalk convention. Control-domain identifiers should be aligned to the Section 4 taxonomy numbering at integration. Note that the relevant adversary technique catalogue is MITRE ATT&CK (Valid Accounts), not ATLAS — a fact the analysis in Section 7.5.6 turns to the framework's advantage.

Table E-2. Failure-to-control-domain crosswalk.

Observed failure	AISIF layer	Control domain	Framework reference
Single-factor authentication on internet-reachable tenants	Operations	Identity trust and authentication	ISO 27001 Annex A identity controls; NIST SP 800-207; platform MFA capability
Credentials unrotated for years; known-compromised credentials valid	Operations	Credential lifecycle and revocation	ISO 27001 Annex A authentication management; infostealer intelligence integration
No trusted-location network policy on tenants	Operations	Network access restriction	ISO 27001 Annex A network controls: zero-trust segmentation
No detection of anomalous authentication or bulk egress	Operations / Data	Behavioral monitoring	ISO 27001 A.8.16; NIST AI RMF Manage; ATT&CK T1078 detections

Contractor endpoints outside customer identity and endpoint standards	Governance	Third-party and supplier control	ISO 27001 supplier-relationship controls; NIST AI RMF Govern third-party actions
Platform-available controls not mandated by customer standards	Governance	Shared-responsibility assignment	CSA AICM role model (AI Customer); ISO/IEC 42001 management clauses
Pipeline/service identities sharing human-credential weaknesses	Application	Non-human identity management	Least-privilege service identities; vaulted short-lived credentials

7.5.4 Maturity-Level Indicators

As with Case Study D, the maturity reading is layer-differentiated — but the profile differs instructively. The affected organization's were not immature technology operators; several run sophisticated security Programs. The Ad Hoc indicators concentrate in one place: identity governance over a third-party data platform. This supports a refinement to the AISIF assessment instrument: maturity must be scored per layer and per estate, because an organization's posture over its own infrastructure and its posture over its SaaS data estate can sit at different levels simultaneously.

Table E-3. Observed maturity indicators by layer (customer-side, affected tenants).

AISIF Layer	Observed indicator	Indicative level
Governance	Platform security capabilities not bound to mandatory standards; contractor access outside customer identity requirements.	Level 1 — Ad Hoc
Data	Analytical estate not treated as crown-jewel data; no egress anomaly detection evidenced by internal discovery.	Level 1–2
Model	Not assessable from the incident; latent exposure only.	Not assessable
Application	Non-human identity posture not directly evidenced; inferred parity with human-credential weaknesses.	Not assessable
Operations	Single-factor authentication; multi-year credential staleness; no network policy; detection by external intelligence.	Level 1 — Ad Hoc

Assessment caveat: indicators are inferred from the joint Mandiant-Snowflake disclosure and characterise the affected tenant configurations at the time of the campaign, not the victim organisations' Programs in their entirety. Approximately 9,000 Snowflake customers were not affected; the indicators describe the compromised cohort's common profile.

7.5.5 Mitigations

Mitigations are stated in AISIF layer order, with the recursive feedback loop as the closing requirement.

- Governance — convert platform-available controls into mandatory customer standards: MFA, network policy and rotation requirements bound to the tenant as a condition of production data residency; extend identity and endpoint requirements contractually and technically to third-party contractors; maintain an inventory of which datasets in the estate feed AI workloads, so data-supply-chain risk is assessable.

- Data — classify the analytical estate at the sensitivity of its most sensitive contents; monitor for anomalous bulk egress with thresholds calibrated to normal pipeline behaviour; treat AI-feeding datasets as a protected category whose compromise triggers Model-layer review.
- Application — issue dedicated, least-privilege, short-lived service identities for AI and analytics pipelines; vault all pipeline credentials; prohibit shared human-service credentials; monitor non-human identity behavior with the same rigor as human accounts.
- Operations — enforce phishing-resistant MFA on every tenant; define maximum credential lifetimes with revocation triggered by infostealer marketplace intelligence; apply trusted-location network policies; detect authentication anomalies (impossible travel, commodity-VPN origin, dormant-account revival) before egress begins.
- Feedback loop — route third-party platform incidents into the governance cycle with authority to amend the customer's own standards: the systemic change here is not a Snowflake configuration fix but a binding shared-responsibility standard applied across every SaaS platform in the data estate. Anticipate replication — the investigation itself warned that the campaign pattern will be repeated against similar SaaS platforms.

7.5.6 Analytical Significance for the Framework

Case Study E closes a coverage argument that Case Study D opened. D showed an AI provider failing through conventional infrastructure controls; E shows AI-consuming enterprises inheriting conventional identity failures through their data supply chain. In both, the adversary technique catalogue that applies is ATT&CK rather than ATLAS, and the preventing controls are ISO 27001 fundamentals — yet the assets at stake are constitutive of the organization's AI capability. The pairing yields the framework's scoping principle: the AISIF Data layer extends upstream to the platforms where AI-relevant data resides, and the Operations layer extends outward to the identities — human and non-human, employee and contractor — that can reach that data.

The case also sharpens the shared-responsibility argument for the standards mapping in Section 5. The CSA AI Controls Matrix assigns control responsibility across provider and customer roles; Case Study E is the empirical demonstration of what happens when that assignment is left implicit. Every preventing control existed; the failure was that no governance instrument made any of them binding on the party positioned to apply them. For practitioners, the operational translation is direct: a control that is available but unmandated should be treated, for risk purposes, as absent.

Sources and Verification Note

Primary sources: Mandiant threat intelligence publication on UNC5537 (cloud.google.com/blog/topics/threat-intelligence, June 2024) and the joint Snowflake-Mandiant-CrowdStrike statements of May-June 2024, including Snowflake's hardening guidance of 30 May 2024. Corroborating reporting: SecurityWeek, Computer Weekly, The Hacker News, Axios and subsequent retrospectives (2024-2026). Victim identifications reflect public confirmations or widely reported attributions as of this writing; the figure of approximately 165 affected organization's is Mandiant's notification count. Where the public record is silent — including the full victim list and per-

victim data scope — this case study states the absence of evidence rather than an inference. Verify particulars against the Mandiant publication before print.